

Ja ihminen loi älyn -podcast

Jakso 2: Neuroverkot olivat suomalaisten juttu – nyt ne pyörittävät tekoälyä – Erkki Oja

Litteraatti:

Erkki Oja:

Minulle vähän naureskeltiin aina, kun kävin Amerikassa jossain workshopissa.

"Mikä neural modeling?"

"Hyi hyi". Vieläkin väitän, että he eivät näe metsää puilta.

Eric Malmi:

Minkä takia neuroverkot toimivat niin hyvin, kuin ne toimivat?

Erkki Oja:

En tiedä. Eikä kukaan tiedä.

Arno Solin:

Tervetuloa Aalto-yliopiston ja Ihminen loi älyn -podcastin pariin! Minun nimeni on Arno Solin, ja toimin täällä Aallossa koneoppimisen professorina ja Suomen ELLIS-instituutissa tutkijana. Mukana minulla on Eric Malmi..

Eric:

Toimin myös Aallossa ja Suomen ELLIS-instituutissa vierailevana professorina.

Arno:

Tänään on kunnia saada vieraaksi suomalaisen tekoälytutkimuksen legenda, emeritusprofessori Erkki Oja. Erkki on tehnyt urauurtavaa työtä neuroverkkojen ja koneoppimisen parissa jo paljon ennen nykyisen tekoälybuumin alkua. Hänet tunnetaan maailmalla muun muassa Ojan säännöstä, joka on nimetty hänen mukaansa. Hän on kouluttanut yli 50 tohtoria, ja meidän nykyinen laitosjohtaja sanoikin kerran, että suomalaisista tekoälyalan professoreista suurin osa on tavalla tai toisella Erkin opetuslapsia. Tervetuloa Erkki Oja!

Erkki:

Kiitos.

Eric:

Erkki, jäit eläkkeelle vuonna 2015. Sen jälkeen on noussut valtava tekoälybuumi. Millä mielin olet seurannut alan kehitystä tässä viimeisen 10 vuoden aikana?

Erkki:

Erittäin hämmästyneenä ja tavallaan myös harmissani, koska todella kävi juuri niin, että viimeinen syksy, kun opetin vielä, siis pidin luentoja, niin yhtäkkiä sellainen tekoälykurssi räjähti. Olikohan siellä 600 osallistujaa tai jotain. Ajattelin, että jos tämä tästä vielä jatkuu, niin tätähän ei pidä hanskassa enää kukaan. Mutta silloin se todella

räjähti melkein samaan aikaan. Vuoden 2015 aikana. Sitten tietysti kun tuli ChatGPT ja nämä, niin se vain kiihtyi. Tavallaan olisi ollut kiva olla silloin mukana, mutta toisaalta oli myös kiva olla silloin, kun sai vähän kehittää sitä alaa kaikessa rauhassa. Eikä ollut sellaista valtavaa opetuspainetta ja ulkopuolelta tulevaa painetta.

Eric:

Muistan sinun, olisiko se ollut jokin tällainen eläkkeelle jäämisen juhlapuheenvuoro. Mainitsit, että sinulla on musta nahkakantinen punareunainen vihko, johon on kirjattu jotain tällaisia tutkimusideoita, joihin sinulla on nyt aikaa tutustua, kun jäät eläkkeelle. Mitä näille kuuluu? Oletko jatkanut tutkimuksia näiden parissa?

Erkki:

Ensinnäkin suosittelisin kaikille nuorille tutkijoille, että kannattaisi olla semmoinen. Onko se nyt sitten musta vahakantinen vihko vai mikä se on, mutta kun saat jonkun idean, niin kirjoittelet sen sinne. Minä olen niitä pikkuisen katsellut, tietysti jälkikäteen joitain onnistunut vähän ratkaisemaan. Ne olivat monesti semmoisia idean alkioita, että tekemällä jonkun ihan Matlab-ohjelman, niin siinä alkoi päästä jo jyvälle. Nyt minulla on vihdoinkin ollut aikaa ohjelmoida, niin olen sitä pikkuisen harrastanut. Mutta ei sieltä nyt enää mitään semmoista käänteentekevää ideaa ole kylläkään tullut.

Arno:

Jos mennään ajassa taaksepäin sinne suomalaisen tekoälyn alkulähteille niin sanotusti, niin jos nyt kirjoitettaisiin suomalaisen tekoälyn historiikki, niin mistä se lähtisi liikkeelle?

Erkki:

Se on vaikea kysymys, kun tämä termi tekoäly sinänsä on elänyt ihan valtavasti tässä vuosikymmenten aikana. Ehkä juuri nyt sitten vihdoinkin voisi sanoa, että tekoälyn historiikki alkaa siitä, missä tämä neuroverkkojen tutkimus alkoi. Mutta joskus 10-20 vuotta sitten sitä ei olisi missään nimessä laskettu tekoälyyn, jolloin sitten olisi pitänyt puhua näistä oikeista tekoälymenetelmistä. Se olisi ollut aivan toisenlainen se historiikki. Mutta jos nyt neuroverkot sitten hyväksytään tekoälyyn kuuluviksi, niin sitten se alkaa jostain sieltä 1970-luvun alusta.

Arno:

Mitä silloin tapahtui?

Erkki:

No kyllä tämä professori Teuvo Kohonen, joka oli siis elektroniikan ja teknillisen fysiikan professori, niin hän kiinnostui nimenomaan, sitä ei edes kutsuttu neuroverkkotutkimukseksi silloin, mutta kiinnostui siitä. Hän oli erittäin syvällisesti perehtynyt siihen, miten tietokone toimii. Hän rakensi oman tietokoneenkin siellä labrassa ja julkaisi kirjan digitaalisista piireistä. Hän rupesi miettimään, että miten kähän ne aivot mahtavat toimia, sellaisella insinööriotteella, tietämättä biologiasta ehkä siinä vaiheessa paljon mitään. Mutta hän rupesi tutkimaan ns.

assosiaatiomuisteja, jotka olisivat jonkinlainen hyvin hyvin karkea malli sille, miten ihmisen muisti toimii. Sieltä alkoivat tulla ensimmäiset julkaisut todella silloin 70-luvun alussa. Siitä näkisin, että nykyisen suomalaisen tekoälyn historiikki suunnilleen alkaa.

Eric:

Mainitsit, että tekoäly- ja neuroverkkotutkimus olivat vähän niin kuin kaksi erillistä asiaa. Voisiko avata vähän, mihin tekoäly viittasi silloin alkuvaiheessa?

Erkki:

Kyllä tekoäly silloin viittasi siihen, mihin se tavallaan viittaa nytkin, eli tällaiseen menetelmäjoukkoon, johon kuuluvat esimerkiksi kielen kääntäminen tai konenäkö, robotiikka, hakualgoritmit, tällaiset asiat. Mutta nimenomaan ei neuroverkot, tai ei neuroverkoilla toteutettuna, vaan muunlaisilla tekniikoilla.

Se tekoälyn ihan syvä ydin, ei varmaan menetelmiltään, mutta tutkimusaiheiltaan, on pysynyt hämmästyttävän muuttumattomana sieltä 50-luvulta asti, jolloin sitä alettiin maailmalla kehittää.

Mutta Suomessa jotenkin se perinteinen tekoäly ei kauhean vahvoilla ollut. En tiedä onko nytkään. Tai on varmaan nyt. En halua kollegoja mitenkään mollata. Mutta nämä neuroverkot olivat silloin Suomessa tavallaan dominoiva asia 70-, 80- ja 90-luvulla. Ja sitten maailma tuli perässä, ja niistä tuli yhtäkkiä dominoiva kansainvälisestäkin joskus tuossa vuosituhaten vaihteen jälkeen. Me olimme tietysti hyvin iloisia siitä.

Arno:

Mutta juuri 70- ja 80-luvulla, niin olitko sinä tekoälytutkija ylipäätään? Oliko se termi, jota käytit, kun esittelit itsesi? Mikä tutkija sinä olit?

Erkki:

En missään nimessä. Minuthan olisi heitetty ulos, jos olisin mennyt jonnekin tekoälykonferenssiin. Että "mitä sinä täällä teet, sinähän olet matemaatikko tai tilastotieteilijä". Ja "neuroverkot ovat ihan naurettavia". "Niistä ei ikinä tule mitään". Se oli sen puolen ihmisten käsitys. Mutta nimenomaan siis olen ollut neuroverkkotutkija. Tietysti sitä voi kutsua myös koneoppimiseksi, koska neuroverkoissahan oleellista ovat ne oppimissäännöt. Eihän siellä mitään verkkoja rakennella, vaan sehän on kaikki softaa. Mutta tietokoneohjelma, joka siis oppii datasta on ollut se keskeinen. Tein väitöskirjan vuonna 1977 tavallaan niistä Kohosen assosiaatiomuistimalleista, ja miten ne voisivat toteutua semmoisessa takaisinkytketyssä neuroverkkopiirissä.

Eric:

Mitkä siihen aikaan olivat sellaiset isot pitkän aikavälin tavoitteet tai tulevaisuuden visiot ehkä sekä teillä neuroverkkotutkijoilla että tekoälytutkijoilla? Oliko teillä samat sovelluskohteet mielessä, vai oliko se jotenkin eriytynyt siinäkin mielessä?

Erkki:

Oli se aika lailla eriytynyt. Perinteisessä tekoälyssä oli silloin vielä semmoinen aikamoinen optimismi vallalla, koska ne onnistuivat tekemään sellaisia pieniä prototyyppisysteemejä. Oli semmoisia niin sanottuja palikkamaailmoja, joissa tekoäly ymmärsi, mikä palikka on toisen päällä ja toisen alla ja miten niitä voisi siirrellä. Sitten tämmöisiä joitain alkeellisia, mikä se on, missä pannaan rukseja ja ympyröitä kolme kertaa kolme ruudukkoon. Tämmöisiä ne osasivat ratkoa ja vastaavia aika pieniä ongelmia.

Silloin varmaan aika moni kuvitteli, sanotaan nyt 70-luvun alussa, että ei tässä tarvitse kuin vähän panna lisää tehoa niihin algoritmeihin ja pikkuisen miettiä miten ne olisivat tehokkaampia, niin kyllä ne pystytään ratkaisemaan. Silloinhan oltiin optimistisia, että joku kahden kielen kääntäminen toisekseen ei voi olla vaikea tie. Se on aika lyhyenkin aikavälin tavoite.

Mutta heillä siis selvästi oli tämmöisiä ihan konkreettisia isoja päämääriä, että nyt tämä, tämä ja tämä tehdään koneälyllä eikä tarvita sitten enää välttämättä ihmisälyä. Eli hyvin samantapaisia toiveita, kuin nykyään on näillä deep learning -tutkijoilla. Kun taas sen neuroverkkopuolen lähtökohdat olivat ihan toisenlaiset. Se oli hyvin vahvasti sidoksissa tämmöiseen neuraaliseen mallitukseen, kognitiotieteeseen. Ei silloin oikeastaan ajateltu, että niillä pystyisi ratkaisemaan jotain, vaikkapa teknisiä ongelmia, vaan yritettiin ihan vilpittömästi miettiä miten ne neuronit toimivat aivoissa. Mitä neuroneista koostuvat verkot oikein tekevät siellä, ja onko siellä joitain sellaisia toimintaperiaatteita, joista voisi oppia jotain? Ja suorastaan ihan konkreettisesti, kun kumminkin oli paljon aivotutkijoita silloin, jotka tekivät suoranaisia mittauksia, upottelivat mikroeletteodeja neurosoluihin ja katsoivat miten ne laukoivat siellä. Niin pystyttäisiinkö ne kokeet mallintamaan? Vähän niin kuin fysiikassa tehdään koe ja yritetään keksiä teoria, miksi se toimii näin. Eli kognitiotiede ja neuromallitus olivat se lähtökohta neuroverkkotutkimukselle.

Eric:

Kohtasivatko nämä kaksi yhteisöä missään konferensseissa? Oliko teidän välillä jotain väittelyä tai sellaista vihanpitoa? Tai millainen se dynamiikka oli?

Erkki:

Eivät ne ainakaan positiivisessa mielessä missään kohdanneet. Parhaimmillaan ne olivat aivan erillään. Kumpikaan ei oikein tiennyt mitä ne muut tekevät. Ehkä eivät olleet kiinnostuneitakaan.

Sitten tietysti neuroverkkotutkijat hiljakseen alkoivat valua sinne käytännön sovellusten puolelle. Voitaisiinko tehdä vaikka joku käsinkirjoitettujen numeroiden tunnistaminen neuroverkoilla, eikä niillä perinteisillä menetelmillä?

Tai vaikka voitaisiinko joku päättely- tai asiantuntijajärjestelmä tehdä neuroverkoilla eikä tämmöisillä sääntöpohjaisilla tai loogisilla kielillä, kuten Lisp ja Prolog? Silloin sitä törmäystä alkoi tulla. Yleensä kumpikin puoli oli vähän pessimistinen sen vastapuolen menetelmillä, että mahtavatko nuo toimia.

Eric:

Ehkä tässä on tapahtunut jonkinlaista konvergoitumista, kun nyt yksi suosituimpia sovelluskohteita näillä syväoppimisjärjestelmillä, neuroverkoilla, on koodin

generoiminen. Eli tavallaan ne generoivat tällaisia symbolisia sääntöpohjaisia järjestelmiä. Siinä mielessä voidaan hyödyntää neuroverkkotutkimusta näiden perinteisten tekoälyn tyyppisten menetelmien kehittämiseen.

Erkki:

Aivan.

Arno:

Tämä perinteinen sääntöpohjaisuus ja erilaisten symbolien manipulointi ja matemaattisten lausekkeiden todistaminen ja muu oli tosiaan pitkään dominoiva näkökulma tekoälyyn. Sitten tämmöinen datalähtöinen oppiminen ylipäättään on ehkä tosiaan paljon uudempi asia sitten taas. Että älyä voikin oppia.

Erkki:

Ilman muuta juuri näin. Silloin ei tietenkään kukaan olisi voinut kuvitellakaan, että neuroverkoilla voisi nyt jotenkin oppia niitä päättelyjärjestelmiä tai mitään tällaista. Oltiin ehkä vähän sokeita sille, koska todellinen syy näiden neuroverkojen tai syväoppimisen nykyiseen menestykseen on tietysti se valtava datan määrä ja sitten se valtava laskentateho.

Silloin oltiin ehkä vähän sokeita sille, että näiden määrä tulee eksponentiaalisesti kasvamaan. Olihan siinä ne Mooren lait ja muut. Silloin kun ei ollut edes mikrotietokoneita ja oli vain muutama vähän isompi tietokone siellä täällä, niin ei tätä raa'an voiman mietetöntä kasvua, joka tässä on tapahtunut, sanotaan viimeisen 50 vuoden aikana, niin ei sitä voinut varmaan kukaan ennustaa. Eikä myöskään sitä, että sillä raa'alla voimalla saa tällaisia tuloksia aikaan. Aina ajateltiin, että kyllä jotain hienompaa täytyy olla siellä menetelmien taustalla.

Arno:

Tässä mainittiinkin jo Teuvo Kohonen, ja voikin varmaan sanoa, että hän on kenties tunnetuin suomalaisen neuroverkkotutkimuksen pioneeri, nyt jo tietysti edesmennyt akateemikko. Hän toimi sinun väitöskirjasi ohjaajana. Mitä kokemuksia työskentelystä hänen alaisuudessaan oli silloin uran alkuvaiheilla? Lähdit kuitenkin tavallaan vähän eri suuntaan tutkimuksen kanssa, kuin mitä Teuvo oli tehnyt.

Erkki:

Joo, sehän meni sillä tavalla, että olin myös ihan Teuvoa tuntematta ollenkaan niin opiskelijana kiinnostunut tämmöisestä hermosolumallituksesta, aivomallituksesta. Diplomityönikin oli nimeltään joku toisen asteen hermosolun malli. En selosta sitä tässä. Ihan tämmöistä biomatematiikkaa, sanoisinko näin. Mietin sitten, että minnekäs minä lähtisin jatko-opiskelemaan, koska näytti siltä, että haluan jatko-opiskelijaksi. Tapasin sitten Teuvan jossain, en muista missä. Hän kysyi minua hänen labraansa. Kun pääsin armeijasta 1974 tammikuussa, niin minä sitten menin sinne Teuvon laboratorioon. Kuten sanottu, niin väittelin tohtoriksi vuonna 1977 ja olin tiiviissä yhteistyössä Teuvon kanssa. Se oli kova koulu. Teuvo oli vielä niitä vanhan ajan hyviä professoreita. Sa labra oli hyvin pieni. Olisiko siellä ollut alle 10 ihmistä. Siellä oli

ehdoton mies- ja ääniperiaate. Teuvo oli se mies, ja hänen oli se ääni. Pojat tekivät mitä Teuvo käski. Mutta siinä oppi paljon.

Eric:

Jos puhutaan alasta yleisemmin, niin tässä on ollut varmaan paljon erilaisia ylä- ja alamäkiä. Puhutaan tällaisista AI winter -jaksoista. Varmaan sekä neuroverkko- että tämän perinteisemmän tekoälytutkimuksen suhteen.

Voitko kertoa tästä? Miltä tämä ala on näyttänyt pitkällä aikajänteellä? Nyt eletään kauheaa buumia. Oliko sellaisia vaiheita jo aikaisemmin urallasi? Miten ne ovat kehittyneet siitä eteenpäin?

Erkki:

Joo, tuo on ihan hyvä. Sekä se, sanotaanko, perinteinen tekoäly, että sitten tämä neuroverkkopuoli, ovat menneet sellaista vuoristorataa. Molemmissa on ollut pari kolme semmoista niin sanottua "winteriä", ja sitten taas niitä huippuja.

Niiden hyvin suuri syy loppujen lopuksi on ollut USA:ssa saatavilla oleva rahoitus. Eli siellä on NSF, ja DARPA, puolustusteollisuuden apurahasäätiö. Sitten kun ne huomasivat, että ei tästä taida mitään tulla, ja rahahanat menivät kiinni, niin se koko ala pakostakin supistui, koska siellähän tyypillisesti proffa rahoittaa jatko-opiskelijansa niillä mainituilla stipendeillä. Sitten, kun taas näytti siltä, että tästä tulee jotain, niin sitä rahaa tuli enemmän.

Tilanne oli todella se, että silloin, no, 60-luvulla olin koulupoika, en tiedä mitä silloin tapahtui. Mutta 70-luvulla.. Ehkä vielä se perinteinen tekoäly, niin kuin sanotaan, "good old-fashioned AI", GOFAI, oli jollain tavalla voimissaan. Mutta nimenomaan tässä neuroverkkopuolella kyllä silloin, kun aloitin, oli hyvin hiljaista. Suomessa onneksi ei oikein tiedetty näistä talvista ja kesistä yhtään mitään. Me vain painoimme, ja Teuvo onnistui sitä rahaa saamaan. Hän painoi sitä tutkimusta eteenpäin.

Minulle aina naureskeltiin, kun kävin Amerikassa jossain workshopissa, että "mikä neural modeling" ja "hyi hyi", "luuletko, että siitä mitään tulee". Syynä siihen, miksei sen neural modeling sitten vetänyt oli se, että jos tarjosi näitä mallejansa niille kokeellisille fysiologeille, niin hehän naureskelivat, että "ei se neuroverkko kuule noin toimi", "ei sinne päinkään". Heillä kokeellisina tutkijoina oli sellainen ajatus, että siinä pitää sitten joka ikisen yksityiskohdan loksauttaa kohdalleen. Jos sieltä joku hormoni tai vastaava puuttui, niin eihän tämä näin voi toimia tämä neuron. Vieläkin väitän, että he eivät näe metsää puilta. Täytyyhän niissä aivoissakin nyt kai joku periaate olla. Voi olla, että ei ole mitään periaatteita. Mutta eihän se nyt voi olla vaan, että ne miljardit neurosolut siellä ovat hyvin monimutkaisia, että jokainen on kuin tehdas. Mallita nyt sitten miljardien tehtaiden verkostoja. Eihän siitä mitään tule. Joka tapauksessa, ei se neural modeling oikein onnistunut.

Ja sitten yksi syy oli vielä se, kun tämä Frank Rosenblatt julkisti semmoisen perceptron-neuroverkkomallin. En nyt muista olisiko se ollut 60-luvun lopulla vai 70-luvun alussa. Sitten taas tällainen MIT:ssä toimiva, en tiedä oliko silloin MIT:ssä, mutta myöhemmin ainakin, tämä Marvin Minsky, perinteisen tekoälyn guru, niin hän lyttysi sen tällaisella kirjalla Perceptrons. Se todella teki huonoa näille.

Olin kerran jossain konferenssissa 80-luvun lopulla, minne Minsky tuli, ja ihan pyyteli anteeksi koko ajan. Että "sori vaan, en tiennyt, että monikerrosverkot näin hyvin toimivat, mutta kyllä minä sittenkin olin oikeassa". Se oli hassua kuunneltavaa. Mutta sillä oli suuri vaikutus.

Mutta sitten tulivat nämä asiantuntijajärjestelmät tuossa 80-luvun alussa. Puhuttiin suureen ääneen, että japanilaisilla on tällainen viidennen sukupolven tietokoneprojekti. Menin Japaniin 1983-1984 Tokion tekniseen korkeakouluun vierailevaksi tutkijaksi ihan sen takia, katsomaan, että mitähän ne oikein tekevät siellä. Siellä valitettavasti ilmeni, että eihän siitä varmaan mitään tule. Mutta siellä nyt oltiin sitten.

Sitten taas neuroverkkopuolella räjähti, kun tämä John Hopfield esitti niin sanotun Hopfield-neuroverkon, olikohan se 1983. Ja Teuvo Kohonen esitteli Self-organizing map -neuroverkon. Muistaakseni 1986 tuli tällainen kuuluisa kirja, Rumelhart ja Williams, olikohan Hintonkin siinä kirjoittajana, kuin Parallel Distributed Processing. Se oli todella merkittävä. Siinä esiteltiin nämä monikerrosverkot. Silloin se lähti vauhtiin.

En tiedä onko tämä perinteinen tekoäly sitten sen 80-luvun puolenvälin jälkeen, en ole hirveästi seurannutkaan, mutta onko se oikein siitä toipunut varsinaisesti ollenkaan. Mutta se neuroverkkobuumi oli aika vahva. Se pikkuisen hiljeni tuossa 90-luvun lopulla, 2000-luvun alussa, kunnes sitten taas tuli tämä Hintonin AlexNet. Semmoinen neuroverkko, jolla tunnistettiin miljoonia kuvia, ja jossa hän esitteli tämän deep learning -tekniikan. Siitä tämä sitten on ollut ihan voittokulkua tähän päivään asti.

Eric:

Kun puhutaan näistä menetelmistä, niin kuin Arno mainitsi tuossa alussa, niin sinut tunnetaan tästä Ojan säännöstä. Voitko avata vähän, että mistä siinä on kyse? Ja miten se liittyy tähän suurempaan kenttään ja näihin muihin menetelmiin? Tai aivojen toimintaan?

Erkki:

No ensinnäkin siitä tunnettuudesta. Olin vierailemassa kerran tuolla Bostonin yliopistossa, jossa on tällainen Stephen Grossberg -niminen iso guru ja se hänen joku jatko-opiskelijansa. Minä kysyin, että tunnetko tällaisen Oja's rule -nimisen säännön. Hän oli, että "tunnen tunnen, kaikkihan sen tuntee". No miksi se tunnetaan? "Katso, kun se Oja on niin lyhyt nimi". Siinä saattaa olla ihan vinha perä.

Tuommoinen Oja's rule jää mieleen. Ei kaikki tiedä mikä se on, mutta joku semmoinen siellä nyt on. Siitä on ollut valtava apu, että sattuu olemaan tällainen lyhyt nimi, josta kukaan ei edes tiedä mitä kieltä se on. Että onko se espanjaa vai mitä se on. Se on siis neuraalinen mallitus. Tarkoituksena oli miettiä, että mitä jos semmoisen hermosolun oikein nyt kutistaa. Sehän on siis elävä solu, se on valtavan monimutkainen. Mutta eihän sille hermosolulle ole oleellista se metabolismi, joka siellä on. Sille on oleellista se, että sinne menee niitä säikeitä sisään, sitten ne tekevät jotain, sitten sieltä tulee säie ulos ja niitä aktiviteetteja. Eli se mallittaa tämän minimissään hyvin semmoisena input/output-tyyppisenä laskentayksikkönä.

Mutta se oppiminen, joka sen aikaisen käsityksen mukaan ja tietääkseni nykyäänkin jopa näiden fysiologien mielestä tarkoittaa sitä, että ne synapsit eli hermosolujen väliset kytkennät muuttuvat. Eli kun ne synapsit saavat uudet arvot, niin silloin aivot toimivat pikkuisen eri tavalla, ja sitä kutsutaan sitten oppimiseksi. Tai älykkyydeksi, miksi sitä kutsutaankin.

Sitä synaptista oppimissääntöä oli jossain määrin tutkittu. Mutta minua se kiinnosti hirveästi. Siis se sääntö tarkoittaa sitä, että jos niillä synapseilla on tietty arvo nyt, ja sitten sinne työnnetään sisään jotain inputtia, niin miten se arvo muuttuu seuraavalla ajanhetkellä. Eli jos nyt sanotaan w :llä niin kuin "weight", niin w_{t+1} on w_t jotain. Näitä oppimislakeja me olimme Kohosenkin kanssa miettineet yhdessä. Sitten vain esitin tällaisen säännön, joka on äärimmäisen yksinkertainen. Hokkus pokkus mitä se tekee on, että se sisään tulevasta signaalivirrasta kaivaa niin sanotut pääkomponentit, eli optimaalisesti kompressoii sen. Että eteenpäin lähtee sellainen kompressoitu esitys siitä.

Jos ajatellaan, että varmaan näiden, ainakin sensoristen neuroverkkojen, siis näön, kuulon, näihin liittyvien neuroverkkojen siellä alemmalla tasolla täytyy olla joku semmoinen systeemi, joka poistaa sitä redundanssia. Ajatellaan, mitä tässäkin huoneessa on. Me luulemme, että tuolla on kattoikkuna ja kirjahylly, mutta täältä tulee mieletön määrä valoa vain meidän silmiin itse asiassa. Jotenkin saadaan sitä hirvittävää informaatio-, data- tai bittimäärää pienennettyä niin, että aivot alkavat tekemään tolkkua siitä, mitä ne oikein näkevät.

Tämä pääkomponenttimenetelmä on yksi kaikkein vahvimpia menetelmiä juuri siinä, miten saadaan redusoitua sitä informaatiota. Se yllättävän paljon herätti kiinnostusta. En oikein ollut arvannutkaan, että niiden fysiologien parissa, kumminkin tämä idea putosi kypsään maahan, niin siihen alkoi sitten tulla viittauksia aika hyvin. Minä en tietenkään kehdannut sitä Ojan säännöksi nimetä, mutta sitten sitä ruvettiin joissain siteeraavissa papereissa kutsumaan nimellä Oja rule. En minä nyt vastustanutkaan. Antaa mennä.

Arno:

Tämähän on kuitenkin erilaista kuin valtavirta tällä hetkellä, jossa sitten taas näiden isojen mallien opetus perustuu backproppiin ja siihen optimointiin ihan eri tavalla. Eikä siihen, että sieltä datasta löydetään pääkomponentteja. Mutta emmehän me tiedä mitä menetelmiä meillä on 10 vuoden kuluttua tai 20 vuoden kuluttua. Selvästi meidän aivomme ovat aika tehokkaita oppijoita. Tuskin siellä mitään backpropagationia tapahtuu sinänsä. Niin ehkä tulevaisuus vielä tuo Ojan säännön keskiöön.

Erkki:

Täsmälleen niin kuin sanoit. Aivoissa ei taatusti voi olla backpropagationia. Se on kyllä ihan totta. Siis se backpropagation on tällainen oppimissääntö näille monikerrosverkoille, jotka syöttävät sitä signaalia eteenpäin koko ajan, ja sitten niillä on semmoinen tavoitesignaali, johon ne haluavat pyrkiä tällä modifikaatiolla. Niissähän tietysti myös niillä jokaisella synapsilla on semmoinen.

Kun kirjoitat sen esille, niin ei se niin kauhean erilainen ole verrattuna siihen Ojan sääntöön. Siinä on vain se ero, että siellä on mukana myös se ylhäältäpäin tuleva, se haluttu outputti. Mutta tietenkään silloin kun tätä kehitin, niin eihän ollut olemassa mitään backpropagationia. Ei voinut verratakaan siihen. Kukaan ei kysynyt, että miksi sinä tuollaista teet, mikset sinä tee backproppia. Backpropagation esiteltiin viisi vuotta sen jälkeen, kun olin sen paperini kirjoittanut. Sitten tässä on semmoinen tärkeä juttu, että...

Eric:

Saanko kysyä tuosta pienen aasinsillan? Kun puhutaan, että backprop esiteltiin, niin tästä on vähän niin kuin käyty kiistaa jossain internetin palstoilla, että kuka sen keksi, ja että onko tässäkin Suomiyhteys. Niin mikä on sinun näkemyksesi?

Erkki:

Ei siinä oikein ole Suomiyhteyttä, kun ei täällä oikeastaan juuri kukaan tutkinut sitä. Teuvo Kohonen, joka oli se suurmies, niin hänellä oli se oma verkkonsa ja hän sanoi, että täällä ei tutkita backproppia hänen labrassaan. Tutkikoot muut sitä. Eli tämä "not invented here", tämä "NIH"-ilmiö. Niin sitten sitä ei täällä myöskään tutkita.

Mutta tätä keskusteluahan on ollut. Voi voi. Nyt kun Geoff Hinton sai Nobel-palkinnon ja ennen sitä sen Turing awardin, niin niitä kateellisia on ollut hirvittävästi. Ja syystäkin, jos on ollut ihan hilkulla, tavallaan toinen keksijä, mutta ei siinä samassa ryhmässä, eikä saa mitään, ja toinen saa nämä isot palkinnot. Kyllähän voisin nimetä muutamia henkilöitä, jotka ovat olleet todella kitkeriä ja lähettelleet nettiin kirjoituksia. Että todistakaa nyt kaikille, että minä sen keksin. Itse asiassa aika huvittava juttu, joka Suomeen liittyy kylläkin on se, että Hinton itsekin siteeraa tätä Seppo Linnainmaata.

Eric:

Tähän juuri viittasin.

Erkki:

Hän vuonna 1970 esitti sen backpropin. Mutta eihän hän neuroverkkoja tutkinut, vaan pyörästysvirheiden etenemistä tämmöisessä verkossa. Mutta tavallaan algoritmisesti kyllähän hänet mainitaan siellä.

Arno:

Palaisin vielä sinne 80-luvulle. Tässä on mainittu paljon semmoisia todella suuria nimiä, jotka yhdistetään nyt tekoälyn voittokulkuun ja historiaan tavallaan jälkeinpäin. Mutta minua näin tekoälytutkijana itse kiinnostaisi myös siirtyä sinne 80-luvun tapaamisiin karpäsenä kattoon kuulemaan mitä se keskustelu on ollut silloin, mitä se interaktio on ollut. Neural information processing systems, joka nykyään identifioidaan varmaan tekoälyn huippufoorumiksi, niin sen juuret ovat siellä jo tämän nykyisen neuroverkkohistorian alkuhämärissä. Onko sinulla mitään anekdoottia jakaa sieltä 80-luvulta?

Erkki:

Joo, ei nyt ehkä anekdoottia. Mutta totta on, että tietysti jokainen ala lähtee sellaisella aika pienellä porukalla. Sanotaan nyt ihan, että tietysti Teuvo, joka oli se todellinen pioneeri, niin hän kuului jo ihan semmoiseen alle kymmeneen ihmiseen, jotka jotenkin nousivat esille sieltä. Niin sitten näiden kymmenen oppilaat hiljaksen myös pääsivät mukaan näihin piireihin. Sitä kautta minäkin sitten. Todella tämä NIPS oli yksi. NIPS workshopit olivat aina jossain hiihtokeskuksessa. No sellainen anekdootti, että Breckenridgen lippis varastettiin minulta Rukalla, kun olin hiihtämässä. Mutta Aspenin lippis minulla vielä on.

Sitten toinen, ehkä vielä merkittävämpi kokous oli tällainen snowboard-workshop. Se oli myös jostain 80-luvun puolivälistä. Se oli kutsupohjainen, ja minä sitten sain onneksi sinne kutsun. Siellä oli joku satakunta ihmistä. Ne olivat hirveän upeita workshoppeja kerta kaikkiaan. Siellä oli paitsi amerikkalaisia niin myös Euroopasta, Japanista ja muualtakin parhaat tutkijat. Se keskustelu oli hyvin intensiivistä.

Eikä siellä tietenkään nyt enää sitten jotain backproppia jauhettu, kun se oli jo keksitty, vaan esim. Geoff Hintonkin esitteli aina vaan uusia ja uusia, Boltzmann machinejä, vaikka mitä. Tavallaan se ura, joka johti sitten tähän nykyiseen, oli ehkä enemmän tuo teollisuus. Ja ehkä jatko-opiskelijat, jotka kaipasivat jotain aiheita. Ne jaksoivat niitä keksittyjä ideoita siellä työstää. Mutta nämä todella uudet ajatukset tulivat usein niistä workshoppeista. Se oli intensiivistä touhua. Jos nyt sitten katsoo, että keitä siellä oli, niin huomaa, että he ovat kyllä juuri näitä kovia nimiä. He ovat nyt sitten joko akateemisessa maailmassa tai sitten näiden isojen firmojen siellä jossain tutkimuspuolella.

Arno:

Meitä kiinnostaa myös muut perinteet tekoälyn älyn parissa ja yksi semmoinen asia, jonka minä ainakin henkilökohtaisesti liitän tosi vahvasti tähän tekoäly-yhteisöön ja nimenomaan näihin koneoppimis- ja neuroverkkoporukoihin, on menetelmien ja koodien avoimuus. Semmoinen ideoiden jakamisen kulttuuri. Tämä tuntuu erottavan tekoälytutkimuksen monista muista aloista, joilla ehkä pidetään kovemmin kiinni niistä omista mittausdatoista, menetelmistä tai muista asioista.

Varsinkin oman urani varrella olen hyötynyt siitä paljon, että muut tutkimusryhmät laittavat koodia, menetelmiä ja myös dataa jakoon. Dataa jaetaan tutkimusryhmien välillä, jotta pystyttäisiin testaamaan uusia menetelmiä. Onko sinulla näkemystä tästä? Onko sinulla samanlaisia kokemuksia uran varrelta? Kuinka uusi ilmiö tällainen on?

Erkki:

En tiedä, jos sanot, että se on nimenomaan neuroverkkoolalle ominaista, niin uskottavahan se on. Mutta kyllähän tämä nyt aika paljon perustutkimuksessa silloin kun ollaan perustutkimusvaiheessa, niin tietojenkäsittelytieteessä varmaan yleisestikin vallitsee. Kyllähän sinne laitetaan niitä koodeja niin kauan kuin niistä ei ole vielä mitään kaupallisia sovelluksia.

Data on tietysti vaikeampi asia siitä syystä, että sen pitää olla oikeaa dataa. Tässä on todella tärkeä seikka, jota esim. Teuvo aina korosti, että älkää leikkikö millään simuloidulla tai keinotekoisesti tuotetulla datalla. Silloin te saatte virheellisiä tuloksia. Sen datan täytyy olla oikeaa reaalimaailman dataa. Ja mistä semmoista saa?

Sehän on ollut koko minunkin tutkimuksessa koko ajan sellainen ongelma, että pitääkö minun juosta tuolla firmoissa rukoilemassa, että antakaa minulle dataa. No nehän tulevat epäluuloisiksi. Eivät ne mitään dataa anna. Yleensä datan omistaa joku. Tietysti nyt on tämä netti. Silloin ei ollut nettiä. Nyt se on, niin voihan sinne vain mennä. Niin kuin nämä isoimmat firmat tuntuvat tekevän, että otetaan sieltä kaikki ja kysytään jälkikäteen, jos silloinkaan. Mutta kyllä siinä yleensä niihin joku copyright kumminkin sisältyy. Se on vaikea asia.

Mutta totta on tietysti, että kun näin dataintensiivistä tietotekniikan alaa ehkä ei ole aikaisemmin ollutkaan, kuin tämä data mining, big data, tekoäly, niin ehkä siellä nyt sitten on muodostunut sellaisia käytäntöjä, että jos joku tutkimusryhmä saa semmoisen datan mitä he saavat levittää, niin kyllähän niitä aika auliisti sitten tietysti muillekin annetaan. Se on totta.

Eric:

Entä ennen netin tuloa, jaettiin niitä datasettejä jo siinä vaiheessa? Lähetettiinkö niitä joillain disketeillä sitten postitse?

Erkki:

Kyllä varmaan jaettiin. Tämä tavallaan jopa ohjasi sitä koko työskentelyä. Esimerkiksi meillä Kohosen labrassa ja minun labrassani myöhemmin harrastettiin tuota kuvan ja videon käsittelyä, koska sinähän voit mennä tuonne kampukselle ja ottaa valokuvan. Se on sitten silloin sinun kuvasi. Eli kuvamateriaali on sellaista, joka ei ole firmojen omistamaa välttämättä. Ainakin kohtuullisissa määrissä sitä voi omasta kotialbumista suunnilleen saada.

Toinen oli sitten puheentunnistus. Myös puheen pölinää saa helposti mistä vain. Mutta semmoiset eksoottisemmat kuva-aineistot, sanotaan vaikka satelliittikuvien käsittely, olivat mielenkiintoisia. Siinä sitten VTT:ltä kaukokartoituslaboratoriosta piti pyytää, että saammeko me tutkimuskäyttöön niitä. Lääketieteelliset kuvat ovat aika hankalia monestakin syystä. Niitä on vaikea saada.

Se tavallaan ohjasi jossain määrin myös ihan sitä tutkimuksen suuntautumista. Menetelmät ovat samoja, mutta kuitenkin ainahan ne täytyy sovittaa tiettyyn dataan. Kyllä tämä datapullonkaula on olemassa. Se on hankala. Siinä on ehkä jopa yksi syy, miksi nykyään tätä tutkimusta tekevät ylivoimaisesti vahvimmin ne isot firmat isolla rahalla. Eivät niinkään yksityiset tutkimuslaboratoriot, niin kuin melkein millä tahansa muulla tieteenalalla.

Eric:

Jos mennään vähän filosofisemmille vesille, niin minkä takia neuroverkot toimivat niin hyvin kuin ne toimivat?

Erkki:

En tiedä. Eikä kukaan tiedä. Se on ihan se iso asia, joka minua askarruttaa. Jos vielä olisi hirveästi halua ja voimaa tehdä tutkimusta, niin yrittäisin mennä sinne. Mutta se on tavattoman vaikea kysymys, että minkä takia.

Te asiantuntijoina tiedätte tämän paremmin kuin minä, mutta kaikki katsojat eivät ehkä tiedä. Siis se monikerrosverkkohan on tällainen neuroverkko, jossa on peräkkäisiä hyvin homogeenisia kerroksia suuri määrä. Jo silloin 80-luvulla todistettiin, vai olisiko ollut 90-luvulla kun todistettiin, että semmoinen verkko, missä on yksi ainoa niin sanottu piilokerros, eli on inputit, piilokerros ja sitten output-kerros, niin se pystyy jo approksimoimaan mielivaltaista jatkuvaa funktiota. Jos on kaksi piilokerrosta, niin sen funktion ei tarvitse olla edes jatkuva.

Jolloin luulisi, että okei, siinähan se on. Sitten vain rakennetaan se verkko, ja me voimme approksimoida. Koska melkein kaikki ne tehtävät, joita näillä neuroverkoilla tehdään, ovat pohjimmiltaan funktion approksimoimista. Tämä vaatii ehkä vähän nielemistä, mutta tämä maailma on täynnä funktioita, joissa inputista tulee outputti. Kun pysähdytte miettimään, niin kaikki on sitä, että inputista tulee outputti. Siinä välillä on niin sanottu funktio. Semmoinen epälineaarinen monimutkainen laskentakaava, joka muuttaa sen inputin outputiksi. Tämän funktion approksimoimisesta on siinä koko hommassa kysymys, vieläpä hyvin perinteisillä menetelmillä. Joskin siellä on paljon sellaisia hienoja yksityiskohtia. Mutta enemmän tai vähemmän regressiotahan se on. Luulisi, ja kaikki luulivat silloin, että ei niitä kauhean montaa kerrosta voi laittaa, koska se tulee liian monimutkaiseksi. Se verkko alkaa oppia sitä opetusdataa, eikä enää yleistä hyvin. Ja naureskeltiin, että hyvänen aika, jos on miljoona parametria niin tarvitaan miljoonan kokoinen opetusdata.

Se on aika huvittavaa, koska esim. World Wide Webin teksti, niin onkohan se, en nyt muista sitä lukua, mutta se on tietysti miljardeja. Sieltä vain röyhkeästi joku Geoff Hinton tai joku sanoo, että okei, otetaan miljardin kokoinen datajoukko ja lyödään se sinne. Sitten se verkkokin saa olla iso. Tämä oli aika vastaansanomaton väite. Tuntuu olevan vielä niin, en nyt seuraa ihan aktiivisesti, mutta jotkut väittävät, että mitä isompi verkko, sitä paremmin se toimii. Se on ehkä aika inkrementaalinen parannus enää siellä, kun mennään vaikka 1000 kerroksesta 2000 kerrokseen. Silti se on parannus. Tilastotieteilijä sanoisi tässä, että ei voi olla.

Tämä on niin syvälinen kysymys. Joko se todistaa jotain tästä tällaisten valtavien funktioiden geometriasta, jota me emme ole ymmärtäneet, kun me olemme tottuneet, että se käyrähän on tällainen x - y . Tässä on x ja y , ja siinä menee tämä käyrä. Se on funktio. Mutta ovatko nämä valtavat funktiot ja vielä strukturoidut funktiot... Nyt minä puhun pitkään. Sopiiko? Että funktio on aina painotettu summa edellisistä funktioista, jotka ovat painotettuja summia edellisistä funktioista. Että onko tässä maailmassa semmoinen hierarkkinen rakenne, että se rakentuu samalla tavalla? Että kaikki koostuu osista, jotka koostuvat osista, jotka koostuvat osista. Jotenkin siihen tämä neuroverkko pääsee pureutumaan.

Eli onko maailma vähän erilainen, kuin me olemme luulleet? Ainakin tämä sensorinen maailma. Vai onko siinä matemaattisessa rakenteessa jotain sellaista, jota matemaatikot eivät ole saaneet purtua selville? Tämä on se, jos minä nyt ryhtyisin neuroverkkotutkijaksi, niin tätä minä rupeaisin tutkimaan. Joskin on siellä niin paljon etevämpiä kavereita sen kimpussa, että tuskin siinä pystyisi mitään tekemään. Mutta jos olisin huippuälykäs neuroverkkotutkija, niin sinne minä nyt menisin. Että miksi monikerrosverkko toimii niin hyvin kuin se toimii?

Eric:

Tästä vinkkiä nuorille tutkijanaluille.

Erkki:

Kyllä. Sinne voi teroittaa kynsiään.

Arno:

Minusta tuntuu, että kaikki huippuasiantuntijat, joiden kanssa olen keskustellut sanovat ihan samaa. Kenelläkään ei ole vastausta tähän. Tämä on kyllä todella syvälinen ja kiinnostava kysymys.

Erkki:

Niin. Joku sanoo, että universumin rakenteen selvittämiseen tarvittiin Einstein, joka keksi suhteellisuusteorian ja nämä aika-avaruusjatkumot. Mutta tämän älykkyyden ratkaisemiseen tarvitaan 5 tai 10 Einsteiniä. Kannattaisi aloittaa.

Arno:

Ehkä meidän pitää vaan treenata tarpeeksi kyvykäs tekoälymalli, joka pystyy tämän meille selittämään.

Erkki:

Hyvä hyvä. Ehkä siinä voisi toimia tällainen vähän niin kuin näissä siruissa. Että tarvitaan tehokkaat sirut kehittämään ne seuraavat sirut. Niin tässäkin tekoäly inkrementaalisesti aina parantaisi itseään. Ehkä sieltä se lopullinen äly sitten löytyisi.

Eric:

Tämän parissa on nyt syntynyt useampia yrityksiä. Tällainen rekursiivinen itsensä parantaminen. Varmasti tullaan näkemään jotain sen saralla.

Erkki:

Siinä on omat vaaransa tietysti myös sitten.

Eric:

Ehkä palataan niihin.

Arno:

Palataan siihen kohtaan.

Sinun tapauksessasi tämä kannatti, tämä tietty jääräpäisyys ja se, että pysyit siellä neuroverkkojen parissa ja loit uran siellä. Vaikka tämä eksponentiaalinen kiinnostuksen kasvu ja tämä tekoälybuumi ajoittuvatkin tavallaan siihen, kun hyppäsit liikkuvan junan kyydistä eläkkeelle.

Siitä huolimatta meitä Ericin kanssa kiinnostaa, että tutkijoina kuinka itsepäisesti kannattaa pitää kiinni siitä omasta tutkimusagendasta ja jatkaa sen saman tutkimusalueen parissa, eikä tavallaan hyppiä vähän aiheesta toiseen sen mukaan, että mikä tuntuu olevan muiden mielestä se oikea polku sillä hetkellä?

Erkki:

Joo. Pikkuisen riippuu ehkä siitäkin, että mihin tähtäät. Totta kai jos nyt joku nuori ihminen haluaisi saada mahdollisimman paljon julkaisuja, saada esimerkiksi jonkun viran yliopistosta, professuurin, voidakseen sitten sen jälkeen mennä vaikka tekemään teollisuutta, niin ihmisellä ei ole välttämättä se tieteellinen intohimo päällimmäisenä. Vaan se, että mikä on optimistrategia sille, että pääsee tieteellisellä uralla eteenpäin. Niin se voikin olla se, että hajautat, teet fiksujen ihmisten kanssa, ja koitat saada nimesi aina niihin papereihin, koskivat ne paperit sitten ihan mitä tahansa. Tämähän on yksi tapa ja kaikki tiedämme ihmisiä, jotka tämmöistä noudattavat.

Mutta toinen on se, että jos nyt ihan oikeasti olet kiinnostunut siitä tutkimusaiheestasi, niin mistä sen sitten tietää, että se oma idea on hyvä. Teuvo, minun oppi-isäni sanoi, että jotenkin sen haistaa. Joko sinä tiedät tai et tiedä. Jos sinä et tiedä, niin olet ehkä huono tutkija. Mutta hyvä tutkija tietää, mikä on arvokasta ja mikä on roskaa. Ja on ihan mahdotonta selittää kenellekään mistä se tulee. Mutta sinä niin kuin tiedät ja uskot asiaasi.

Tähän tietysti tulee sellaisia viitteitä sitten, että esimerkiksi konferensseissa joku edes on kiinnostunut, tai sitten sinun papereihisi alkaa tulla viittauksia. Mutta kyllä semmoinen idea tai vähintään kapea tutkimusala pitää olla selvillä, ja sitten sitä täytyy vaan jyyttää. Jos jätät sen kesken, niin siihen se jää.

Esimerkiksi minulla tämä niin sanottu Oja's rule. Kehitin sen noin vuonna 1980. Siitä tuli sitten tämmöinen uusi tekniikka nimeltään riippumattomien komponenttien analyysi, josta minulla oli varmaan viimeiset paperit joskus 30 vuotta myöhemmin. Aina vaan siitä samasta erilaisilla muunnelmilla. Pitää olla sinnikäs, todella sinnikäs. Ei se aina onnistu, mutta se on yksi tapa.

Arno:

Olemme keskustelleet isoista nimistä, ja tunnet henkilökohtaisesti tai olet työskennellyt monien maailman johtavien tekoälytutkijoiden, kuten Teuvo Kohosen tai Nobel-palkitun Geoffrey Hintonin kanssa. Onko jotain sellaisia erityisiä piirteitä tai kykyjä, jotka erottavat nämä ihmiset muista?

Erkki:

Ehkä sellainen luonteenpiirre, että eivät usein ole kovin mukavia ihmisiä. Siis se on varmasti ihan kaikilla aloilla. He jotka pääsevät huipulle ovat särmikkäitä, tämmöisiä Elon Muskeja. Niin kuin amerikkalaiset sanovat, "nice guys finish last". Se on aika hyvin sanottu. He ovat särmikkäitä tyyppejä, joiden varpaille ei astuta. Ei mitään kovin kilttejä ihmisiä.

Ja sitten sinnikkäitä ja kauhean älykkäitä. Ihan mielettömän fiksuja. Kyllä siellä takki korvissa hiiviskelee sieltä workshopista hotellihuoneeseen. Että miten nuo voivat olla tuommoisia? Siellä tulee semmoinen alemmuudentunto suorastaan, kun juttelee jonkun aivan tällaisen huipputyyppin kanssa. He tietävät hirveästi, ja heillä leikkaa aivan hirvittävän hyvin. Kyllä siinä ihmisten fiksuudessaakin eroja on.

Eric:

Ovatko nämä sellaisia synnynnäisiä luonteenpiirteitä enemmän, vai voiko siinä kehittää itseään?

Erkki:

Pelkään, että ei.

Arno:

Minä toivon, että ei.

Erkki:

Kyllä ihminen on semmoinen, miksi se syntyy. Pilata voi mutta ei oikein positiiviseen suuntaan kauheasti voi.

Eric:

Tämä nyt ehkä vastaa tai jotenkin tyhjentää tämän seuraavan kysymyksen. Mutta onko sinulla jotain neuvoja, joita jättäisit nuorille tutkijoille? Tai ehkä sellaisille, jotka ovat kiinnostuneita alasta, eivät ole välttämättä tutkijoita, mutta ajattelevat, että tämä tekoäly on nyt kiinnostava juttu?

Erkki:

Nimenomaan juuri tuohon äskeiseenkin viitaten se, että pitäisi aina hakeutua itseään fiksumpien seuraan. Sehän on se konsti. Että ei missään nimessä ikinä saa mennä sellaiseen ryhmään, josta heti näet, että minähän olen fiksumpi kuin nämä. Tietysti siellä voi ehkä loistaa, mutta ei siinä ole mitään järkeä. Pitää valita mahdollisimman fiksut kollaboraattorit, jotka tekevät kanssasi töitä.

Minulla oli onni olla Teuvon kanssa, joka on tällainen huippututkija. Toinen oli tämä Leon Cooper, jonka mainitsin, Nobelfyysikko. Hänen labrassaan myös todella näki, mitä se suuren maailman tutkimus on. Että opi viisaammilta.

Arno:

Joissain piireissä ollaan vähän varuillaan tekoälyn nopean kehityksen kanssa, ja puhutaan sellaisesta käsitteestä kuin P(doom). Eli siitä, kuinka monen prosentin todennäköisyydellä tekoälyn uskotaan johtavan ihmiskunnan tuhoon. Mikä sinun P(doom) todennäköisyytesi on?

Erkki:

En ollenkaan uskonut tuohon vielä ehkä viisi vuotta sittenkään. Tai 10 vuotta sitten ainakaan. Se tuntui aivan hölynpölyltä. Kaikki nämä singulariteetit ja muut. Mutta nyt sitten muutama todella minua fiksumpi, esimerkiksi Geoff Hinton, tai Hawkins, joka nyt ei ehkä tunne alaa, mutta kuitenkin on tästä varoittanut ilmeisen vilpittömällä mielellä, että tällainen vaara on olemassa. Niin en minä nyt nollana pidä sitä todennäköisyyttä missään nimessä enää. Ehkä nämä kaverit, jotka ovat pitkään alalla olleet ja nähneet kaiken, niin niillä saattaa olla joku intuitio tai sisäinen käsitys siitä, että tämä ei mene hyvin.

Hinton esimerkiksi, minä todella kunnioitan häntä, hänen nimensä nousee monesti esille, on käynyt täällä minun vieraananikin. Hän näki vielä sen suurteollisuuden ja sen, kuinka ahneita ne tyypit ovat ja piittaamattomia tämmöisistä turvallisuustekijöistä, koska silloin pikkuisen menetetään rahaa, kun pannaan rajoituksia sinne. Niin sanotaan, että ei ole mahdotonta, että tämmöinenkin vaara voisi olla olemassa. Tietysti oikein pitkällä ajanjaksolla se on hyvinkin todennäköistä, mutta ei nyt tässä 10-20 vuoteen.

Arno:

Kiitoksia.

Eric:

Suuret kiitokset.

Erkki:

Oliko se tässä?

Eric:

Se oli tässä.

Erkki:

Ei kysymystä tietoisuudesta?